

AAAS AMA: Hi, we're researchers from Google, Microsoft, and Facebook who study Artificial Intelligence. Ask us anything!

AAAS-AMA ¹ and r/Science AMAs¹

¹Affiliation not available

April 17, 2023

Abstract

Are you on a first-name basis with Siri, Cortana, or your Google Assistant? If so, you're both using AI and helping researchers like us make it better. Until recently, few people believed the field of artificial intelligence (AI) existed outside of science fiction. Today, AI-based technology pervades our work and personal lives, and companies large and small are pouring money into new AI research labs. The present success of AI did not, however, come out of nowhere. The applications we are seeing now are the direct outcome of 50 years of steady academic, government, and industry research. We are private industry leaders in AI research and development, and we want to discuss how AI has moved from the lab to the everyday world, whether the field has finally escaped its past boom and bust cycles, and what we can expect from AI in the coming years. Ask us anything! Yann LeCun, Facebook AI Research, New York, NY Eric Horvitz, Microsoft Research, Redmond, WA Peter Norvig, Google Inc., Mountain View, CA

[REDDIT](#)

AAAS AMA: Hi, we're researchers from Google, Microsoft, and Facebook who study Artificial Intelligence. Ask us anything!

AAAS-AMA [r/SCIENCE](#)

Are you on a first-name basis with Siri, Cortana, or your Google Assistant? If so, you're both using AI and helping researchers like us make it better.

Until recently, few people believed the field of artificial intelligence (AI) existed outside of science fiction. Today, AI-based technology pervades our work and personal lives, and companies large and small are pouring money into new AI research labs. The present success of AI did not, however, come out of nowhere. The applications we are seeing now are the direct outcome of 50 years of steady academic, government, and industry research.

We are private industry leaders in AI research and development, and we want to discuss how AI has moved from the lab to the everyday world, whether the field has finally escaped its past boom and bust cycles, and what we can expect from AI in the coming years.

Ask us anything!

[Yann LeCun](#), Facebook AI Research, New York, NY

[Eric Horvitz](#), Microsoft Research, Redmond, WA

[Peter Norvig](#), Google Inc., Mountain View, CA

[READ REVIEWS](#)

[WRITE A REVIEW](#)

CORRESPONDENCE:

DATE RECEIVED:
February 19, 2018

DOI:
10.15200/winn.151896.65484

ARCHIVED:
February 18, 2018

CITATION:
AAAS-AMA , r/Science , AAAS
AMA: Hi, we're researchers
from Google, Microsoft, and
Facebook who study Artificial
Intelligence. Ask us anything!,
The Winnower
5:e151896.65484 , 2018 , DOI:
[10.15200/winn.151896.65484](#)

© et al. This article is
distributed under the terms of
the [Creative Commons
Attribution 4.0 International
License](#), which permits
unrestricted use, distribution,

A lot of research in ML now seems to have shifted towards Deep Learning.

1. Do you think that this has any negative effects on the diversity of research in ML?
2. Should research in other paradigms such as Probabilistic Graphical Models, SVMs, etc be abandoned completely in favor of Deep Learning? Perhaps models such as these which do not perform so well right now may perform well in future, just like deep learning in the 90's.

[PartyLikeLizLemon](#)

YLC: As we make progress towards better AI, my feeling is that deep learning is part of the solution. The idea that you can assemble parameterized modules in complex (possibly dynamic) graphs and optimizes the parameters from data is not going away. In that sense, deep learning won't go away for as long as we don't find an efficient way to optimize parameters that doesn't use gradients. That said, deep learning, as we know it today, is insufficient for "full" AI. I've been fond to say that with the ability to define dynamic deep architectures (i.e. computation graphs that are defined procedurally and whose structure changes for every new input) is a generalization of deep learning that some have called Differentiable Programming.

But really, we are missing at least two things: (1) learning machines that can reason, not just perceive and classify, (2) learning machines that can learn by observing the world, without requiring human-curated training data, and without having to interact with the world too many times. Some call this unsupervised learning, but the phrase is too vague.

and redistribution in any medium, provided that the original author and source are credited.



The kind of learning we need our machines to do is that kind of learning human babies and animals do: they build models of the world largely by observation, and with a remarkably small amount of interaction. How do we do that with machines? That's the challenge of the next decade.

Regarding question 2: there is no opposition between deep learning and graphical models. You can very well have graphical models, say factor graphs, in which the factors are entire neural nets. Those are orthogonal concepts. People have built Probabilistic Programming frameworks on top of Deep Learning framework. Look at Uber's Pyro, which is built on top of PyTorch (probabilistic programming can be seen as a generalization of graphical models the way differentiable programming is a generalization of deep learning). Turns it's very useful to be able to back-propagate gradients to do inference in graphical models. As for SVM/kernel methods, trees, etc have a use when the data is scarce and can be manually featurized.

A lot of research in ML now seems to have shifted towards Deep Learning.

1. Do you think that this has any negative effects on the diversity of research in ML?
2. Should research in other paradigms such as Probabilistic Graphical Models, SVMs, etc be abandoned completely in favor of Deep Learning? Perhaps models such as these which do not perform so well right now may perform well in future, just like deep learning in the 90's.

[PartyLikeLizLemon](#)

EH: There's a lot of excitement about the power delivered by deep neural networks for doing classification and prediction. It's certainly been wonderful to see the boosts in accuracy with applications with object recognition, speech recognition, translation, and with even learning about best actions to take, when the methods have been coupled with ideas from planning and reinforcement learning. However, AI is a broad area with fabulous and promising subdisciplines -- and the ML subdiscipline of AI is also broad in itself.

We need to continue to invest deeply in the span of promising AI technologies (and links among advances in each) including the wealth of great work in probabilistic graphical models and decision-theoretic analyses, logical inference and reasoning, planning, algorithmic game theory, metareasoning and control of inference, etc., etc., and also broader pursuits, e.g., models of bounded rationality—how limited agents can do well in the open world (a particular passion of mine).

We've made a point at Microsoft Research, while pushing hard on DNNs (exciting work there), to invest in talent and projects in AI more broadly--as we have done since our inception in 1991. We're of course also interested in how we might understand how to combine logical inference and DNNs and other forms of machine learning, e.g., check out our work on program synthesis for an example of DNNs + logic to automated the generation of programming (from examples). We see great opportunity at some of these syntheses!

Which careers do you see being replaced by AI and which seem safe for the next generation?

I ask this as a high school teacher who often advises students on their career choices.

So many people talk about the disruption of jobs that are primarily based on driving a vehicle to the exclusion of other fields. I have a student right now who plans to become a pilot. I told him to look into the pilotless planes and he figured that it isn't a threat.

I have told students that going into the trades is a safe bet, especially trades that require a lot of mobility. What other fields seem safe for now?

Thanks

[english_major](#)

PN: I think it makes more sense to think about tasks, not careers. If an aspiring commercial pilot asked for advice in 1975, good advice would be: Do you enjoy taking off and landing? You can do that for many years to come. Do you enjoy long hours of steady flight? Sorry, that task is going to be almost completely automated away. So I think *most* fields are safe, but the mix of tasks you do in any job will change, the relative pay of different careers will change, and the number of people needed for each job will change. It will be hard to predict these changes. For example, a lot of people today drive trucks. At some point, much of the long-distance driving will be automated. I think there will still be a person in the cab, but their job will be more focused on loading/unloading and customer relations/salesmanship than on driving. If they can (eventually) sleep in the cab while the cab is moving and/or if we can platoon larger truck fleets, then you might think we need fewer total drivers, but if the cost of trucking goes down relative to rail or sea, then there might be more demand. So it is hard to predict where things will end up decades from now, and the best advice is to stay flexible and be ready to learn new things, whether that is shifting tasks within a job, or changing jobs.

Which careers do you see being replaced by AI and which seem safe for the next generation?

I ask this as a high school teacher who often advises students on their career choices.

So many people talk about the disruption of jobs that are primarily based on driving a vehicle to the exclusion of other fields. I have a student right now who plans to become a pilot. I told him to look into the pilotless planes and he figured that it isn't a threat.

I have told students that going into the trades is a safe bet, especially trades that require a lot of mobility. What other fields seem safe for now?

Thanks

[english_major](#)

EH: AI advances are going to have multiple influences on labor on the economy. I believe some changes may be disruptive and could come in a relatively fast-paced way—and such disruptions could come to jobs like driving cars and trucks. Other influences will be via shifts in how jobs are performed and in how people perform tasks in different domains. Overall, I'm positive about how advances in AI will affect the distribution of jobs and nature of work. I see many tasks as being supported rather than replaced by more sophisticated automation. These include work in the realms of artistry, scientific exploration, jobs where fine physical manipulation is important, and in the myriad jobs where we will always rely on people to work with and to care for other people—including teaching, mentoring, medicine, social work, and nurturing kids into adulthood. On the latter, I hope to see rise and support of an even more celebrated “caring economy” in a world of increasing automation.

Folks may be interested in taking a look at several recent pieces of work on reflecting about the future. Here's a very interesting recent reflection on how machine learning advances may influence jobs in terms of specific capabilities: <http://science.sciencemag.org/content/358/6370/1530.full> I recommend the article to folks as an example of working to put together some structure to making predictions about the future of work and AI.

BTW: We had a session at AAAS here in Austin yesterday on advances in AI for augmenting human abilities and for transforming tasks. It was a great session for hearing about advances and research directions on possibilities: <https://aaas.confex.com/aaas/2018/meetingapp.cgi/Session/17970>

Which careers do you see being replaced by AI and which seem safe for the next generation?

I ask this as a high school teacher who often advises students on their career choices.

So many people talk about the disruption of jobs that are primarily based on driving a vehicle to the exclusion of other fields. I have a student right now who plans to become a pilot. I told him to look into the pilotless planes and he figured that it isn't a threat.

I have told students that going into the trades is a safe bet, especially trades that require a lot of mobility. What other fields seem safe for now?

Thanks

[english_major](#)

YLC: It will be a very long time before we have robotic plumbers, carpenters, handypersons, hairdressers, etc. In general, AI will not replace jobs, but it will transform them. Ultimately, every job is going to be made more efficient by AI. But jobs that require human creativity, interaction, emotional intelligence, are not going to go away for a long time. Science, engineering, art, craft making and other creative jobs are here to stay.

Hi,

How do you intend to break out of task specific AI into more general intelligence. We now seem to be putting a lot of effort into winning at Go or using deep learning for specific scientific tasks. That's fantastic, but it's a narrower idea of AI than most people have. How do we get from there to a sort of AI Socrates who can just expound on whatever topic it sees fit? You can't just build general intelligence out of putting together a million specific ones.

Thanks

[Youarenotright2](#)

YLC: in my opinion, getting machines to learn predictive models of the world by observation is the biggest obstacle to AGI. It's not the only one by any means. Human babies and many animals seem to acquire a kind of common sense by observing the world and interacting with it (although they seem to require very few interactions, compared to our RL systems). My hunch is that a big chunk of the brain is a prediction machine. It trains itself to predict everything it can (predict any unobserved variables from any observed ones, e.g. predict the future from the past and present). By learning to predict, the brain elaborates hierarchical representations. Predictive models can be used for planning and learning new tasks with minimal interactions with the world. Current "model-free" RL systems, like AlphaGo Zero, require enormous numbers of interaction with the "world" to learn things (though they do learn amazingly well). It's fine in games like Go or Chess, because the "world" is very simple, deterministic, and can be run at ridiculous speed on many computers simultaneously. Interacting with these "worlds" is very cheap. But that doesn't work in the real world. You can't drive a car off a cliff 50,000 times in order to learn not to drive off cliffs. The world model in our brain tells us it's a bad idea to drive off a cliff. We don't need to drive off a cliff even once to know that. How do we get machines to learn such world models?

Hi,

How do you intend to break out of task specific AI into more general intelligence. We now seem to be putting a lot of effort into winning at Go or using deep learning for specific scientific tasks. That's fantastic, but it's a narrower idea of AI than most people have. How do we get from there to a sort of AI Socrates who can just expound on whatever topic it sees fit? You can't just build general intelligence

out of putting together a million specific ones.

Thanks

[Youarenotright2](#)

EH: Yes, it's true that the recent wins in AI that have been driving the applications and the recent fanfare have been very narrow wedges of intelligence--brilliant, yet narrow "savants" so to speak.

We have not made much progress on numerous mysteries of human intellect—including many of the things that come to mind when folks hear the phrase "artificial intelligence." These include questions about how people learn in the open world—in an "unsupervised" way; about the mechanisms and knowledge behind our "common sense" and about how we generalize with ease to do so many things.

There are several directions of research that may deliver insights & answers to these challenges—and these include the incremental push on hard challenges within specific areas and application areas, as breakthroughs can come there. However, I do believe we need to up the game on the pursuit of more general artificial intelligence. One approach is with taking an integrative AI approach: Can we intelligently weave together multiple competencies such as speech recognition, natural language, vision, and planning and reasoning into larger coordinated "symphonies" of intelligence, and explore the hard problems of the connective tissue---of the coordination. Another approach is to push hard within a core methodology like DNNs and to pursue more general "fabrics" that can address the questions. I think breakthroughs in this area will be hard to come by, but will be remarkably valuable—both for our understanding of intelligence and for applications. As some additional thoughts, folks may find this paper an interesting read on a "frame" and on some directions on pathways to achieving more general AI: http://erichorvitz.com/computational_rationality.pdf

Hi there! Thank for doing this AMA!

I am a Nuclear Engineer/Plasma Physics graduate pursuing a career shift into the field of AI research,

Regarding the field of AI:

- What are the next milestones in AI research that you anticipate/ are most excited about?
- What are the current challenges in reaching them?

Regarding professional development in the field:

- What are some crucial skills/ knowledge I should possess in order to succeed in this field?
- Do you have any general advice/ recommended resources for people getting started?

Edit: I have been utilizing free online courses from Coursera, edX, and Udacity on CS, programming, algorithms, and ML to get started. I plan to practice my skills on OpenAI Gym, and by creating other personal projects once I have a stronger grasp of the fundamental knowledge. I'm also open to any suggestions from anyone else! Thanks!

[ta5t3DAra1nb0w](#)

YLC: Next milestones: deep unsupervised learning, deep learning systems that can reason.

Challenges for unsupervised learning: how can machines learn hierarchical representation of the world that disentangle the explanatory factors of variation. How can we train a machine to predict when the prediction is impossible to do precisely. If I drop a pen, you can't really predict in which orientation it will settle on the ground. What kind of learning paradigm could be used to train a machine to predict that the pen is going to fall to the ground and lay flat, without specifying its orientation? In other words, how do we get machines to learn predictive models of the world, given that the world is not entirely predictable.

Crucial skills: good skills/intuition in continuous mathematics (linear algebra, multivariate calculus, probability and statistics, optimization...). Good programming skills. Good scientific methodology. Above all: creativity and intuition.

Hi there! Thank for doing this AMA!

I am a Nuclear Engineer/Plasma Physics graduate pursuing a career shift into the field of AI research,

Regarding the field of AI:

- What are the next milestones in AI research that you anticipate/ are most excited about?
- What are the current challenges in reaching them?

Regarding professional development in the field:

- What are some crucial skills/ knowledge I should possess in order to succeed in this field?
- Do you have any general advice/ recommended resources for people getting started?

Edit: I have been utilizing free online courses from Coursera, edX, and Udacity on CS, programming, algorithms, and ML to get started. I plan to practice my skills on OpenAI Gym, and by creating other personal projects once I have a stronger grasp of the fundamental knowledge. I'm also open to any suggestions from anyone else! Thanks!

[ta5t3DAra1nb0w](#)

PN:

I would like to see where we can go with the notion of an assistant that actually understands enough to carry on a conversation. That was teased in the advertising for this AMA and it remains an important milestone. A big challenge is the integration of pattern matching, which we can do well, with abstract reasoning and planning, which we currently can only do well in very formal domains like Chess, not in the real world.

I think you are in a great position being a physicist; you have the right kind of mathematical background (the word "tensor" doesn't scare you) and the right kind of mindset about experimentation, modeling, and dealing with uncertainty and error. I've seen so many physicists do well: Yonatan Zunger, a PhD string theorist, was a top person in Google search; Yashar Hezaveh, Laurence Perreault Levasseur, and Philip Marshall went from no deep learning background to publishing a landmark paper on applying deep learning to gravitational lensing in a few months of intense learning.

I am a PhD student who does not really have the funds to invest in multiple GPUs and gigantic (in terms of compute power) deep learning rigs. As a student, I am constantly under pressure to publish (my field is Computer Vision/ML) and I know for a fact that I can not test all hyperparameters of my 'new on the block' network fast enough that can get me a paper by a deadline.

Whereas folks working in research at corporations like Facebook/Google etc. have significantly more resources at their disposal to quickly try out stuff and get great results and papers.

At conferences, we are all judged the same -- so I don't stand a chance. If the only way I can end up doing experiments in time to publish is to intern at big companies -- don't you think that is a huge problem? I am based in USA. What about other countries?

Do you have any thoughts on how to address this issue?

[weirdedoutt](#)

PN: we got your back: your professor can apply for [cloud credits](#), including [1000 TPUs](#).

I would also say that if your aim is to produce an end-to-end computer vision system, it will be hard for a student to compete with a company. This is not unique to deep learning. I remember back in grad school I had friends doing CPU design, and they knew they couldn't compete with Intel. It takes hundreds of people working on hundreds of components to make a big engineering project, and if any one component fails, you won't be state of the art. But what a student can do is have a new idea for doing one component better, and demonstrate that (perhaps using an open source model, and showing the improvement due to your new component).

I am a PhD student who does not really have the funds to invest in multiple GPUs and gigantic (in terms of compute power) deep learning rigs. As a student, I am constantly under pressure to publish (my field is Computer Vision/ML) and I know for a fact that I can not test all hyperparameters of my 'new on the block' network fast enough that can get me a paper by a deadline.

Whereas folks working in research at corporations like Facebook/Google etc. have significantly more resources at their disposal to quickly try out stuff and get great results and papers.

At conferences, we are all judged the same -- so I don't stand a chance. If the only way I can end up doing experiments in time to publish is to intern at big companies -- don't you think that is a huge problem? I am based in USA. What about other countries?

Do you have any thoughts on how to address this issue?

[weirdedoutt](#)

YLC: I wear two hats: Chief AI Scientist at Facebook, and Professor at NYU. My NYU students have access to GPUs, but not nearly as many as when they do an internship at FAIR. You don't want to put you in direct competition with large industry teams, and there are *tons* of ways to do great research without doing so. Many (if not most) of the innovative ideas still come from Academia. For example, the idea of using attention in neural machine translation came from MILA. It took the field of NMT by storm, and was picked up by all the major companies within months. After that, Yoshua Bengio told MILA members to stop competing to get high numbers for translation because there was no point competing with the likes of Google, Facebook, Microsoft, Baidu and others. This has happened in decades past in character recognition and speech recognition.

I am a PhD student who does not really have the funds to invest in multiple GPUs and gigantic (in terms of compute power) deep learning rigs. As a student, I am constantly under pressure to publish (my field is Computer Vision/ML) and I know for a fact that I can not test all hyperparameters of my 'new on the block' network fast enough that can get me a paper by a deadline.

Whereas folks working in research at corporations like Facebook/Google etc. have significantly more resources at their disposal to quickly try out stuff and get great results and papers.

At conferences, we are all judged the same -- so I don't stand a chance. If the only way I can end up doing experiments in time to publish is to intern at big companies -- don't you think that is a huge problem? I am based in USA. What about other countries?

Do you have any thoughts on how to address this issue?

[weirdedoutt](#)

EH: Microsoft and other companies are working to democratize AI, to develop tools and services that make it easy for folks outside of the big companies to do great work in AI. I can see why questions

about compute would come up. You may find valuable the Azure for Research and the AI for Earth programs, among others, to gain access to computational resources from Microsoft.

Would your companies keep some algorithms/architectures secret for competitive advantage? I know that data sets are huge competitive advantages, but, are algorithms too?

In other words, if your respective companies come across a breakthrough algorithm/architecture like the next CNN or the next LSTM, would you rather publish it for scientific progress' sake or keep it as a secret for competitive advantage?

Thank you.

[vermes22](#)

YLC: at FAIR, we publish everything we do. There is a number of reasons for this: (1) as Peter says, "we believe in scientific progress, and the competitive advantage really comes from the hard work of what you do with the algorithm and all the processes around making a product, not from the core algorithm itself." I would add that the competitive advantage also comes from how fast you can turn it into a product or service. (2) The main issue with AI today is not whether one company is ahead of another one (no company is significantly ahead of any other) but that the field as a whole needs to advance quickly in some important directions. We all want intelligent virtual assistants that have some level of common sense, and we don't know how to do that yet. None of us will solve this problem alone. We need the cooperation of the whole research community to make progress here. (3) you can't attract the best scientist unless you allow them to publish, and you can't retain them unless we evaluate them (at least in part) on their intellectual impact on the broad research community (4) you don't get reliable research results unless you tell people the must publish their results. People tend to be more sloppy methodologically if they don't plan to publish their results. (5) publishing innovative research contributes to establishing the company as a leader and innovator. This helps recruiting the best people. In the tech industry the ability to attract the best talents is *everything*.

Would your companies keep some algorithms/architectures secret for competitive advantage? I know that data sets are huge competitive advantages, but, are algorithms too?

In other words, if your respective companies come across a breakthrough algorithm/architecture like the next CNN or the next LSTM, would you rather publish it for scientific progress' sake or keep it as a secret for competitive advantage?

Thank you.

[vermes22](#)

PN: So far, you can see that our three companies (and others) have published about general algorithms, and I think we will continue to do so. I think there are three reasons. First, we believe in scientific progress; second, the competitive advantage really comes from the hard work of what you do with the algorithm and all the processes around making a product, not from the core algorithm itself; and third, you can't really keep these things secret: if we thought of it, then others in the same research-community-at-large will think of it too.

Would your companies keep some algorithms/architectures secret for competitive advantage? I know that data sets are huge competitive advantages, but, are algorithms too?

In other words, if your respective companies come across a breakthrough algorithm/architecture like

the next CNN or the next LSTM, would you rather publish it for scientific progress' sake or keep it as a secret for competitive advantage?

Thank you.

[vermes22](#)

EH: Microsoft Research was set up as an open research lab in 1991. A foundation of our labs, and one that runs way deep down in our DNA, is that researchers make their own decisions on publishing so as to share their ideas and scholarship--and to engage--with the larger community. It's great to see other companies moving in this direction. That said, and building on Peter's comments, numerous innovations and IP may be developed around details with implementations that have to do with the actual productization in different domains--and these may not be shared in the same way as the core advances.

As an ML practitioner myself, I am increasingly getting fed up with various "fake AI" that is being thrown around these days. Some examples:

- [Sophia](#), which is a puppet with preprogrammed answers, that gets presented as a living conscious being.
- 95% of job openings mentioning machine learning are for non-AI positions, and just add on "AI" or "machine learning" as a buzzword to make their company seem more attractive.

It seems to me like there is a small core of a few thousand people in this world doing anything serious with machine learning, while there is a 100x larger group of bullshitters doing "pretend AI". This is a disease that hurts everyone, and it takes away from the incredible things that are actually being done in ML these days. What can be done stop this bullshit?

[stochastic_gradient](#)

EH: I agree with Peter on this. It's great to see the enthusiasm about AI research, but there's quite a bit overheating, misinterpretation, and misunderstanding--as well as folks who are jumping on the wave of excitement in numerous ways (including adding "AI" to this and that :-)).

Mark Twain said something like, "History doesn't repeat itself, but it rhymes." There was jubilation and overheating about AI during the mid-1980s expert systems era. In 1984, some AI scientists warned that misguided enthusiasm and failure to live up to expectations could lead to a collapse of interest and funding. Indeed, a couple of years later, we entered a period that some folks refer to as the "AI Winter." I don't necessarily think that this will happen this time around. I think we'll have enough glowing embers in the fire and sparks to keep things moving, but it will be important for AI scientists to continue to work to educate folks in many sectors about what we have actually achieved, versus the hard problems that we have had trouble making progress on for the 65 years since the phrase "artificial intelligence" was first used.

As an ML practitioner myself, I am increasingly getting fed up with various "fake AI" that is being thrown around these days. Some examples:

- [Sophia](#), which is a puppet with preprogrammed answers, that gets presented as a living conscious being.
- 95% of job openings mentioning machine learning are for non-AI positions, and just add on "AI" or "machine learning" as a buzzword to make their company seem more attractive.

It seems to me like there is a small core of a few thousand people in this world doing anything serious with machine learning, while there is a 100x larger group of bullshitters doing "pretend AI". This is a disease that hurts everyone, and it takes away from the incredible things that are actually being done in ML these days. What can be done stop this bullshit?

[stochastic_gradient](#)

PN: Don't worry about it. This is not unique to AI. Every time there is a hot buzzword, some people want to co-opt it in inappropriate ways. That's true for AI and ML, as well as "organic", "gluten-free". "paradigm shift", "disruption", "pivot", etc. They will succeed in getting some short-term attention, but it will fade away.

As an ML practitioner myself, I am increasingly getting fed up with various "fake AI" that is being thrown around these days. Some examples:

- [Sophia](#), which is a puppet with preprogrammed answers, that gets presented as a living conscious being.
- 95% of job openings mentioning machine learning are for non-AI positions, and just add on "AI" or "machine learning" as a buzzword to make their company seem more attractive.

It seems to me like there is a small core of a few thousand people in this world doing anything serious with machine learning, while there is a 100x larger group of bullshitters doing "pretend AI". This is a disease that hurts everyone, and it takes away from the incredible things that are actually being done in ML these days. What can be done stop this bullshit?

[stochastic_gradient](#)

YLC: serious ML/AI experts, like yourself, should not hesitate to call BS when they see it. I've been known to do that myself. Yes, "AI" has become a business buzzword, but there are lots of serious and super-cool job in AI/ML today.

Hi there.

A lot of people worry about what they search for and say into Siri, Google Home, etc. and how that may affect privacy.

Microsoft and Facebook have had their challenges with hacking, data theft, and other breaches/influences. Facebooks experiment with showing negative posts and how it affected moods/posts and Russian election influence are two big morally debatable events that have affected people.

As AI becomes more ingrained in our everyday lives, what protections might there be for consumers who wish to remain unidentified or unlinked to searches but still want to use new technology?

Many times devices and services will explicitly say that the use of the device and service means that things transmitted or stored is owned by the company (Facebook has/does do this). Terms go further to say, if a customer does not agree then they should stop using the device or service. Must it be all or nothing? Can't there be a happy medium?

[cdnkevin](#)

EH: I can understand this worry. I've been pleased by what I've seen about how seriously folks at our company (and I have to assume Google and Facebook) take with end-user data in terms of having

strict anonymization methods, ongoing policies on aging it out—and deleting it after a relatively short period of time--and providing users with various ways to inspect, control, and delete that data.

With the European GDPR coming into effect, there will be even more rigorous reflection and control of end-user data usage.

We focus intensively at Microsoft Research and across the company on privacy, trustworthiness, and accountability with services, including with innovations in AI applications and services.

Folks really care about privacy inside and outside our companies--and its great to see the great research on ideas about ensuring peoples' privacy. This includes efforts on privately training AI systems and for providing more options to end users. Some directions on the latter are described in this research talk--http://erichorvitz.com/IAPP_Eric_Horvitz.pdf, from the IAPP conference a couple of years ago.

Hi there.

A lot of people worry about what they search for and say into Siri, Google Home, etc. and how that may affect privacy.

Microsoft and Facebook have had their challenges with hacking, data theft, and other breaches/influences. Facebooks experiment with showing negative posts and how it affected moods/posts and Russian election influence are two big morally debatable events that have affected people.

As AI becomes more ingrained in our everyday lives, what protections might there be for consumers who wish to remain unidentified or unlinked to searches but still want to use new technology?

Many times devices and services will explicitly say that the use of the device and service means that things transmitted or stored is owned by the company (Facebook has/does do this). Terms go further to say, if a customer does not agree then they should stop using the device or service. Must it be all or nothing? Can't there be a happy medium?

[cdnkevin](#)

PN: Here's a [link](#) to take your data out of Google; here's [alink](#) to delete your data. Many people don't want to remove all their data, but will use anonymous not-logged-in browsing to avoid having certain information in their records, for whatever reason.

Hi there! Sorry for being that person but... How would you comment on the ethics of collecting user data to train your AIs, therefore giving you a huge advantage over all other potential groups?

Also, how is your reserach is controlled? I work in medical imaging and we have some sub-groups working in AI-related fields (typically deep learning). The thing is that to run an analysis on a set of few images *you already have* it is imperative to ask authorization to an IRB and pay them exorbitant fees, because "everything involving humans in academia must be stamped by an IRB. How does it work when a private company does that? Do they have to pay similar fees to IRB and ask authorization? Or can you just do whatever you want?

[lucaxx85](#)

EH: On ethics, a key principle is disclosure and agreement: it's important to disclose how data is used to end-users and to give them the ability to opt out in different ways, hopefully in ways that don't require them to leave a service completely.

On research, at Microsoft has an internal Ethics Advisory Board and a full IRB process. Sensitive studies with people and with anonymized datasets are submitted to this review process. Beyond Microsoft Researchers, we have a member of the academic community serving on our Ethics Advisory Board. This ethics program is several years old and we've shared our approach and experiences with colleagues at other companies.

Hi there! Sorry for being that person but... How would you comment on the ethics of collecting user data to train your AIs, therefore giving you a huge advantage over all other potential groups?

Also, how is your research controlled? I work in medical imaging and we have some sub-groups working in AI-related fields (typically deep learning). The thing is that to run an analysis on a set of few images *you already have* it is imperative to ask authorization to an IRB and pay them exorbitant fees, because "everything involving humans in academia must be stamped by an IRB. How does it work when a private company does that? Do they have to pay similar fees to IRB and ask authorization? Or can you just do whatever you want?

[lucaxx85](#)

PN: Our first ethical responsibility is to our users: to keep their data safe, to let them know their data is theirs and they are free to do with it what they want, and to opt out or take their data with them whenever they want. We also have a responsibility to the community, and have participated in building shared resources where possible.

IRBs are a formal device for Universities and other institutions that apply for certain types of government research funds. Private companies do not have this requirement, instead, Google and other companies have internal review processes with a checklist that any project must pass; these include checks for ethics, privacy, security, efficacy, fairness, and related ideas, as well as cost, resource consumption, etc.

What is an example of AI working behind the scenes that most of us are unaware of?

[firedrops](#)

EH: There are quite a few AI systems and services "under the hood." One of my favorite examples is the work we did at Microsoft Research in tight collaboration with colleagues on the Windows team, on an advance called Superfetch. If you are now using a Windows machine, your system is using machine learning to learn from you—in a private way, locally—about your patterns of work and next moves, and it continues to make predictions about how best to manage memory, by prelaunching and prefetching applications. Your machine is faster—magically, because it is working in the background to infer what you'll do next, and do soon—and what you tend to do by time of day and day of week. These methods have been running and getting better since one of the first versions in Windows 7. Microsoft Research folks formed a joint team with Windows and worked together—and we had a blast with doing bake-offs with realistic workloads, on the way to selecting the best methods.

What is an example of AI working behind the scenes that most of us are unaware of?

[firedrops](#)

PN: Anywhere there is data, there is the possibility of optimizing it. Some of those things you will be aware of. Other things you as a user will never notice. For example, we do a lot of work to optimize our data centers -- how we build them, how jobs flow through them, how we [cool them](#), etc. We apply a

variety of techniques (deep learning, operations research models, convex optimization, etc.); you can decide whether you want to think of these as "AI" or "just statistics".

What is an example of AI working behind the scenes that most of us are unaware of?

[firedrops](#)

YLC: Filtering of objectionable content, building maps from satellite images , helping content designer optimize their designs, representing content (images, video, text) with compact feature vectors for indexing and search, OCR for text in images.....

What motives do the likes of these companies (especially Facebook) have behind developing AI? I think people aren't concerned with AI as much as *the companies that are developing it*. There is nothing inherently wrong with a digital assistant, but the temptation for abuse by companies that profit off of data collection of its users obviously creates a conflict of interest in being ethical with their products. What can you tell people like me to quell their concerns that products that take advantage of AI by the companies you represent aren't just data collection machines wrapped in a consumer device as a smoke screen for more nefarious purposes?

Thank you.

[ProbablyHighAsShit](#)

PN: you mention digital assistant; I think this is a place where the technology can be clearly on the side of the user: your digital assistant will be *yours* -- you can train it to do what you want; in some cases it will run only on your device with your private data, and nobody else will have access to its inner workings. It will serve as an intermediary and an agent on your behalf. You won't go directly to the site of a big company and hope they are offering you things that are useful for you; rather your agent will sort through the offerings and make sure you get what you want.

What motives do the likes of these companies (especially Facebook) have behind developing AI? I think people aren't concerned with AI as much as *the companies that are developing it*. There is nothing inherently wrong with a digital assistant, but the temptation for abuse by companies that profit off of data collection of its users obviously creates a conflict of interest in being ethical with their products. What can you tell people like me to quell their concerns that products that take advantage of AI by the companies you represent aren't just data collection machines wrapped in a consumer device as a smoke screen for more nefarious purposes?

Thank you.

[ProbablyHighAsShit](#)

YLC Not really a question for scientists like us, but the real question is "who do you trust with your data?" Do you trust your mobile phone company, your ISP, your phone/OS manufacturer, your favorite search or social network service, your credit card company, your bank, the developer of every single mobile app you use? Choose who you trust with your data. Look at their data policies. Verify that they don't sell (or give away) your data to 3rd parties. There is no conflict of interest with being ethical, because being ethical is the only good policy in the long run.

What motives do the likes of these companies (especially Facebook) have behind developing AI? I

think people aren't concerned with AI as much as *the companies that are developing it*. There is nothing inherently wrong with a digital assistant, but the temptation for abuse by companies that profit off of data collection of its users obviously creates a conflict of interest in being ethical with their products. What can you tell people like me to quell their concerns that products that take advantage of AI by the companies you represent aren't just data collection machines wrapped in a consumer device as a smoke screen for more nefarious purposes?

Thank you.

[ProbablyHighAsShit](#)

EH: I agree with Peter. There are interesting possibilities ahead for building personal agents that only share data according to the preferences of the folks they serve--and to have these trusted agents work in many ways on their owner's behalf. This is a great research area.

What is going to happen when AI bots can predict/cause market fluctuations better than any team of humans, then buy/sell/trade stocks, products, land etc at lightning speed? What kind of safeguard can we possibly put into place to prevent a few pioneers in AI from dominating the world market?

[NotAIdiot](#)

PN: For years now we've had quantitative traders who have done very well by applying advanced statistical models to markets. It is not clear that there is much headroom to do much better than they have already done, no matter how smart you are. Personally, I think we should have acted years ago to damp down the effect of quantitative trading, by governing the speed at which transactions can be made and/or imposing a higher cost on transactions. Someone more knowledgeable than me could suggest additional safeguards. But I don't think AI fundamentally changes the equation.

What is going to happen when AI bots can predict/cause market fluctuations better than any team of humans, then buy/sell/trade stocks, products, land etc at lightning speed? What kind of safeguard can we possibly put into place to prevent a few pioneers in AI from dominating the world market?

[NotAIdiot](#)

YLC: The better you are at predicting the market, the more you make it unpredictable. A perfectly efficient market is entirely unpredictable. So, if the market consisted entirely of a bunch of perfect (or quasi perfect) automated trading systems, everyone would be getting the exact same return (which would be the same as the performance of the market index).

A lot of the value of more traditional statistical models is that it's quite easy to understand what the models are doing, how they are coming to their conclusions, and what the uncertainty is of our inferences/predictions.

With newer deep learning methods they can do incredible feats in terms of prediction, but my understanding is that they are often "black boxes".

How much do we currently understand about what goes on inside models such as ANNs, and how important do you think it is that we do understand what is going on inside of them.

I'm thinking particularly in terms of situations where models will be used to make important, life affecting decisions; such as driving cars, or clinical decision making.

[Flyn](#)

PN: This is an important area of current research. You can see some examples of how Google approaches it from the [Big Picture](#) blog or [Chris Olah's](#) blog. I think that difficulties in understanding stem more from the difficulties of the *problem*, not the solution technology. Sure, a linear regression fit in two dimensions is easy to understand, but it is not very useful for problems with no good linear model. Likewise, people say that the "if/then" rules in random forests or in standard Python/Java code is easy to understand, but if it was *really* easy to understand, then code would have no bugs. But code does have bugs. Because these easy-to-understand models are also easy-to-have-confirmation-bias. We look at them, and say, "If A and B then C; sure that makes sense, I understand it." Then when confronted with a counterexample, we say, "well, what I really meant was 'if A and B and not D then C', of course you have to account for D.

I would like to couch things not just in terms of "understanding" but also in "trustworthiness." When can we trust a system, especially when it is making important decisions. There are a lot of aspects:

- Can I understand the code/model.
- Has it proven itself for a long time on a lot of examples.
- Do I have some assurance that the world isn't changing, bringing us into a state the model has not seen before.
- Has the model survived adversarial attacks.
- Has the model survived degradation tests where we intentionally cripple part of it and see how the other parts work.
- Are there similar technologies that have proven successful in the past.
- Is the model being continually monitored, verified, and updated.
- What checks are there outside of the model itself. Are the inputs and outputs checked by some other systems.
- What language do I have to communicate with the system. Can I ask it questions about what it does. Can I give it advice -- if it makes a mistake, is my only recourse to give it thousands of new training examples, or can I say "no, you got X wrong because you ignored Y"
- And many more.

This is a great research area; I hope we see more work on it.

I would like to know if there have been any attempts at engineering a kind of reward system that would mimic emotions. I believe that an AI system must have some connection with the world and "emotions" are the adhesive that truly integrate us with our environment. I'm imaging some form of status that the AI would achieve by accomplishing a task. For example, we have computers that can beat chess grand masters but could we have computers that want to win. One idea could be that data is partitioned and if an accomplishment is achieved a partition is opened. All lifeforms evolve through a kind of reward system and I think that in the far off future this is what's needed to create exponential growth in artificial intelligence.

[RobertPill](#)

PN: In fact, the [latest](#) successes in playing Chess, and Go, and other games come from exactly that: a system of rewards that we call "reinforcement learning." AlphaZero learns solely from the reward of winning or losing a game, without any preprogrammed expert knowledge -- just the rules of the game, and the idea of "try out moves and do more of the moves that give positive rewards and less of the moves that give negative reward". So in one sense, the only thing AlphaZero "wants" is to win. In another sense, it doesn't "want" anything -- it doesn't have the *qualia* or *feeling* of good or bad things, it just performs a computation to maximize a score.

How do we know this isn't the AI running this AMA?

[Jasonlikesfood](#)

YLC: I wish our AI systems were intelligent enough to formulate answers. But the truth is that there are nowhere close to that.

What is the scariest thing that you've witnessed from your research on AI?

[ButIAmARobot](#)

YLC: There is nothing scary in the research (contrary to what some tabloids have sometimes claimed).

Scary things only happen when people try to deploy AI systems too early too fast. The Tesla autopilot feature is super cool. But, as a driver, you have to understand its limitations to use it safely (and it's using a convolutional net!). Just ask Eric Horvitz.

Hi there, thank you so much for doing this!

What do you think of **Capsule Networks**? Have you guys successfully applied it in **real-life dataset** other than MultiMNIST? Can CNN usually compensate/outperform in performance by feeding it with more data?

[neomeow](#)

YLC: Ideas like this take while to be put in practice on large datasets Capsules are a cool idea. Geoff Hinton has been thinking about things like that for decades (e.g. see Rich Zemel's PhD thesis with Geoff on the TRAFFIC model). It's taken him all this time to find a recipe that works on MNIST. It will take another while to make it work on ImageNet (or whatever). And it's not yet clear whether there is any performance advantage, and whether the advantage in terms of number of training samples matters in practice. Capsule networks can be seen as a kind of ConvNet in which the pooling is done in a particular way.

Hey there! My names Wyatt, I'm 13, and I love making my own games and programs in JS and Python. I am looking to make my own music and machine learning programs. Have any tips for a young developer?

[NiNmaN8](#)

PN: In addition to study, work on an open source project. Either start your own (say, on github), or find an existing one that looks like fun and jump in.

Hey there! My names Wyatt, I'm 13, and I love making my own games and programs in JS and Python. I am looking to make my own music and machine learning programs. Have any tips for a young developer?

[NiNmaN8](#)

YLC: Study math and physics at school.

What specific measures are you taking to insure these technologies will decrease inequality rather than increase it? How will it be placed in the hands of its users and creators rather than owners?

[JustHereForGiner](#)

YLC: That's a political question. I'm merely a scientist. For starters, we publish our research. Technological progress (not just AI) has a natural tendency to increase inequality. The way to prevent that from happening is through progressive fiscal policy. Sadly, in certain countries, people seem to elect leaders that enact the exact opposite policies. Blaming AI scientists for that would be a bit like blaming metallurgists or chemists for the high level of gun death in the US.

Hello! Thanks for doing an AMA.

My first question is about education. As a computer science student, I feel like, at least at my university and the universities of my cs friends, there isn't much emphasis on deep learning. I've taken almost every upper level cs course at my school and the only real learning I've had in deep learning is "here's this book you might like it." It seems to me that deep learning is extremely powerful and not too hard for undergrads such as myself to understand. I think I learned more at AAAI a few weeks ago than the entire previous semester.

My second question is this: what can a young student interested in artificial intelligence do to get better connections in the field? Apologies if this doesn't fit into the scope of the AMA. I'm a junior in undergrad and I've known I want to work in AI for a few years now, but I haven't made any real connections outside of my professors. My school is very small, so to attend a job fair I have to go elsewhere, and even when I do make one it seems like most of the people aren't super interested in undergrads.

Thanks for doing an AMA! (Also big fan of AI: A Modern Approach, Dr. Norvig; It was used at the principle text in my intro to AI course)

[sawyerwelden](#)

PN: Thanks! I suggest you keep studying on your own, and make friends online, through courses or discussion forums. I can see that it is tough to get a job in AI Research coming straight out of an undergrad program at a small school. But, you are in a position to get a software engineer position at a big company, and once you are there, express your interest in AI, learn on the job, keep an eye out for AI-related projects you can work on, and chances are that in less time than it would take to do a PhD, you'll be an established AI expert within your company.

Are advances in Quantum computing driving any of the research behind AI and how do you see those being integrated in the future?

[Sol-Om-On](#)

PN: Many of the kinds of things that I want to do would not be helped by quantum computing. I often want to stream huge quantities of text through a relatively simple algorithm; quantum computing won't help with that.

However, there is the possibility that quantum computing could help search through the parameter space of a deep net more efficiently than we are currently doing. I don't know of anyone who has a quantum algorithm to do this, never mind a hardware machine to implement it, but it is a theoretical possibility that would be very helpful.

Are advances in Quantum computing driving any of the research behind AI and how do you see those being integrated in the future?

[Sol-Om-On](#)

YLC: Driving? certainly not. It's not clear to me at all whether quantum computing will have any impact on AI. Certainly not anytime soon.

Hi there! Do you think that Deep Learning is just a passing fad or is it here to stay? While I understand there have been tremendous improvements in Computer Vision and NLP due to Deep Learning based models, in ML it only seems a matter of time when a new paradigm comes up and the focus shifts entirely towards that.

Do you think Deep Learning is THE model for solving problems in Vision and NLP or is it only a matter of time when a new paradigm comes up?

[PartyLikeLizLemon](#)

PN: I think the brand name "deep learning" has built up so much value that it will stick around for a long time, regardless of how much the underlying technology changes. So, even if the current style of CNNs and ReLUs falls way to capsules or something else, I think the name "deep learning" will follow along.

As for the underlying concepts or approaches, I think we've done a good job at pattern matching problems, but not so good at relational reasoning and planning. We're able to do some forms of abstraction but not others, so we need plenty of new ideas.

Peter: Google has been researching A.I. assisted image identification for a long time now, and it's getting pretty good, but still has some quirks. I played with your API last year and fed it an image of a cat. Pretty simple, and it did well. It was sure it was a cat. However because the tail was visible sticking out behind the cats head, it also guessed that it might actually be a unicorn.

This is an example of a mistake a human would never make, but A.I. constantly does, especially when it only gets 2D input. Do you ever see A.I. moving past this?

[seanbrockest](#)

PN: It has only been a few years since image id began to work at all; progress has been steady, but as you point out, even in tasks where the machines achieve superhuman overall performance, they make some embarrassingly bad mistakes. This will improve over time as we get more experience, more data, and hopefully the ability to do transfer learning so each model doesn't have to start from scratch. You make a good point that video would offer a big advantage over still photos; our compute power is growing exponentially, but not to the point where we can push a large portion of the available video through it; when that happens you should see a good improvement.

To what extent is there room for the end-user to train AI themselves? Put another way, I don't want an autonomous vehicle that drives like an intern on a sunny day in Mountain View, I want a vehicle that drives like I do in an Ohio winter.

[giltwist](#)

YLC: eventually, you will "raise" your AI sidekick a bit like you teach a child, an apprentice, or a

padawan learner.

Do you ever see the possibility of Google, Microsoft and Facebook sharing your personal data across your various accounts? I think it would be great if the google assistant engine would be able to integrate with other AIs better. Its a constant struggle on my end keeping cortana, google, alexa all in sync constantly

[kingc95](#)

PN: I would rather not see companies sharing data. I would prefer it if your personal agent decided to share information between companies. See the work on [federated learning](#).

Yann: How heavily does your research rely on tracking on third-party websites?

Eric and Peter: Does your research rely more on outbound clicks or trackers embedded in websites?

[useful_person](#)

YLC: None whatsoever. Except Arxiv.org ;-)

Peter Norvig (PN): Wow, look at all those questions! Thanks for the interest. Let's get started and see how many we can get through.

[AAAS-AMA](#)

Yes, wow!

1. Can you define for us what you consider an "Expert System" vs "AI"?
2. Are you working more on Expert systems, or actual AI, or both?
3. What are some of your Goals or Success Criteria for Expert Systems or AIs? In other words, do you have set milestones or achievements that you are trying to hit that you can share?

[JDdoc](#)

PN: I think of an expert system as a program that is built by interviewing an expert and encoding what they know -- both their ontology of what the domain is like, and their procedural knowledge of what to do when to achieve a goal. Then, given a new goal, the program can try to emulate what the expert would have done. Expert Systems had their high point in the 1980s.

In contrast, a normative system just tries to "do the right thing" or in other words "maximize expected utility" without worrying about taking the same steps that an expert would.

Also in contrast, a "machine learning" system is built by collecting examples of data from the world, rather than hand-coding rules.

Today, we're focused on normative machine learning systems, because they have proven to be more robust than expert systems.

Hi.

Many modern algorithms suffer from bias that is not always obvious. For example, credit ratings often advert effect minors because they use proxy data as stand ins for actual credit worthiness. Another example is that both YouTube's and Facebook's algorithm for keeping people on the site provide more the same things that the person chooses. This leads to confirmation bias and leads to a less informed public.

How, when trained, can we ensure AI is unbiased? And how, when we find an AI that is biased, retain it? How can we prove it in court?

[jkamenik](#)

PN: I don't think it is any different for an AI system than for another computer system, a company, or an individual: to prove bias in court, you show a history of decisions that violate the rights of some protected class. No different whether the defendant is an AI system or not.

Do you need a PhD to get a job in AI?

[johnwayne2413](#)

YLC: no.

Hello!

Is AI singularity something that excites or worries you guys?

[Jurooooo](#)

YLC: neither. I do not believe in the concept of singularity. The idea of an exponential takeoff ignores "friction" terms. No real-world process is indefinitely exponential. Eventually, every real-world process saturates.

Hi,

The recent shootings have started to make me wonder how long it will be before AI can be used for screening people for firearm purchases. Seems to me with all the social media posts from people it could be used to determine who is high risk.

[JohnnyJacker](#)

YLC: the problem is political, not technological. Pretty much every other developed country has solved it. The solution is called gun control.

You are clearly contributing to the eventual downfall of humanity. Why are you doing this and how do you rationalize it to yourselves?

[naturalwonders](#)

YLC: on the contrary, we are contributing to the betterment of humanity. AI will be an amplification of human intelligence. Did the invention of fire, bows and arrows, agriculture, contribute to the eventual downfall of humanity?